



ТЕХНИЧЕСКИ УНИВЕРСИТЕТ - СОФИЯ
Факултет Компютърни системи и технологии
Катедра Информационни технологии в индустрията

маг. инж. Ваня Димитрова Иванова

**Откриване на IoT базирани ботнет атаки чрез анализ на
мрежови трафик**

АВТОРЕФЕРАТ

на дисертация за присъждане на образователна и научна степен
„ДОКТОР”

Област на висше образование: 5. Технически науки
Професионално направление: 5.3. Комуникационна и компютърна техника
Научна специалност: Автоматизирани системи за обработка на информация
и управление

Научни ръководители: **проф. д-р инж. Ташо Ангелов Ташев**
 доц. д-р инж. Иво Руменов Драганов

СОФИЯ, 2022

Дисертационният труд е обсъден и насочен за защита от Научния съвет на Катедра Информационни технологии в индустрията към Факултет Компютърни системи и технологии при Технически Университет – София на редовно заседание, проведено на XX.XX.2022 г. (протокол № XX).

Публичната защита на дисертационния труд ще се състои на XX.XX.2022 г. от XX.00 часа в Конферентната зала на БИЦ на Технически Университет – София на открито заседание на научното жури, определено със заповед № XX на Ректора на ТУ – София в състав:

1. Проф. д-р инж. Румен Иванов Трифонов, ТУ-София – Председател
2. Доц. д-р инж. Иво Руменов Драганов, ТУ-София – Секретар
3. Проф. д-р Атанас Велков Атанасов – ХТМУ
4. Чл.-кор. Проф. д-р Петко Христов Петков, пенсионер
5. Доц. д-р Димитър Иванов Пилев, ХТМУ

Рецензенти:

1. Проф. д-р инж. Румен Иванов Трифонов, ТУ-София
2. Проф. д-р Атанас Велков Атанасов – ХТМУ

Материалите по защитата са на разположение на интересуващите се в Катедра Информационни технологии в индустрията, блок 1, стая 1326

Дисертантът е задочен докторант към Катедра Информационни технологии в индустрията на Факултет Компютърни системи и технологии. Изследванията по дисертационния труд са направени от автора.

Автор: маг. инж. Ваня Иванова

Заглавие: „Откриване на IoT базирани ботнет атаки чрез анализ на мрежови трафик“

Тираж: 30 броя

Отпечатано в ИПК на Технически Университет – София

ОБЩА ХАРАКТЕРИСТИКА НА ДИСЕРТАЦИОННИЯ ТРУД

Актуалност на проблема

Съвременната комуникационна инфраструктура се характеризира с огромен мащаб на своето покритие, разнообразие от предавани данни и видове потребители, които я използват за обмен на изключително богато съдържание, както и за предлагане и потребление на услуги. Все по-широката диверсификация от услуги и способности за тяхното предлагане, по жичен и безжичен път, прави последните все по-уязвими при целенасочени недобронамерени действия от страна на потребители, които целят привеждане на поддържащия хардуер и софтуер в състояние, при което съответната услуга да не е достъпна за използване от регулярно използващите я потребители. Подобен тип зловредни въздействия се експлоатират успешно с използването на компрометирани устройства от вида Интернет на нещата (IoT – Internet of Things), за повечето от които е характерна ниската степен на информационна защита, изключително широкото им разпространение в световен мащаб, както в индустриална и публична среда, така и в дома. Ботнет-базираните разпределени атаки са едно от основните средства за блокиране на достъп до услуги през последните повече от десет години и интересът към възможни начини за ограничаване на въздействието им расте.

Цел на дисертационния труд

Целта на дисертационния труд е да се предложат оптимизирани структури на единични и свързани класификатори за откриване и определяне на вида на IoT базирани ботнет атаки, които са по-точни от вече съществуващи класификатори в практиката при еквивалентна или по-ниска изчислителна сложност на реализацията.

За постигане на целта на дисертационния труд е необходимо да се решат следните задачи:

1. Разработване и оптимизация на детектор на IoT базирани ботнет атаки чрез анализ на мрежови трафик, основан на еднослойна невронна мрежа с обратно разпространение на грешката.
2. Разработване и оптимизация на класификатор на IoT базирани ботнет атаки чрез анализ на мрежови трафик, основан на многослойна невронна мрежа с обратно разпространение на грешката.
3. Разработване и оптимизация на детектор и класификатор на IoT базирани ботнет атаки чрез анализ на мрежови трафик, основани на машина с поддържащи вектори.

4. Разработване и оптимизация на детектор и класификатор на IoT базирани ботнет атаки чрез анализ на мрежови трафик, основани на ансамбъл от случайно генерирани дървета за вземане на решения.

5. Разработване и оптимизация на свързани детектор и класификатор от единични класификатори на IoT базирани ботнет атаки чрез анализ на мрежови трафик, които осигуряват по-висока точност на детекция и разпознаване на провежданите атаки в сравнение с точността на изграждащите ги звена.

Научна новост – Разработени са две и три компонентни свързани класификатори за откриване на мрежови атаки, реализирани с помощта на ботнет, изградени от IoT и други видове устройства в пакетно комутируеми мрежи чрез анализ на реализирания трафик към атакувани машини след агрегиране на записите от изградените връзки към тях под формата на обобщени признаци. Представените реализации се състоят от оптимизирани единични класификатори – невронна мрежа с обратно разпространение на грешката, машина с поддържащи вектори и класификатор от типа „случайна гора“. Единичните и съставните класификатори са изследвани по отношение на постигана точност и бързодействие, които показват тяхната приложимост. Видът на откриваните атаки включва DoS TCP, DoS UDP, DoS HTTP, DDoS TCP, DDoS UDP, DDoS HTTP, Keylogging, Data Exfiltration, OS Fingerprint и Service Scan.

Практическа приложимост - Разработените класификатори за откриване на мрежово базирани атаки могат да намерят място в системите за мониторинг на мрежовия трафик с цел превенция на злонамерени въздействия като: централизирано и разпределено влияние с цел отказ от услуга, прехващане на информация, въвеждана от устройства за стандартен вход в компютърните системи, неоторизиран достъп до записана информация в дискови масиви, регистриране на характеристики на операционната система на атакуваната машина, както и на активните услуги, предлагани от нея.

Апробации – Към момента няма внедряване на разработените класификатори в масовата практика.

Публикации - Основните резултати от изследванията са представени в научните публикации, описани в края на автореферата.

Структура и обем на дисертационния труд

Дисертационният труд е с обем от 171 страници и включва съдържание, списък с използваните съкращения, увод, пет глави, списък с научно-приложни и приложни приноси и списък с използвана литература. Номерата на фигурите и таблиците в автореферата съответстват на тези в дисертационния труд.

СЪДЪРЖАНИЕ НА ДИСЕРТАЦИОННИЯ ТРУД

ПЪРВА ГЛАВА - АНАЛИТИЧЕН ОБЗОР НА ИОТ БОТНЕТ

Табл. 1.1. Основни характеристики на най-популярните ботнет

Ботнет	Вид код	Период на откриване	Общ брой засегнати устройства	Уязвими архитектури
Mirai [27, 28, 29, 30, 31, 32, 33, 34, 35, 113]	Само-разпространяващ се червей	Август, 2016	> 600 000	ARC, RCE, MIPS, PowerPC, x86
Dark Nexus, [60, 61, 62, 63]	-	Декември, 2019	Първоначално 1400	12 types - ARM7, etc.
Noaxcalls, [64, 65, 66, 67]	-	Април, 2020	Първоначално около 1300	x86, i686, m68k, PPC, ARM5-7, MIPS, etc.
LeetHozer, [68, 69]	-	Март, 2020	-	i586, rxrg, ARM6 & 7, MIPS, MP5L, SH4, M68K, SPC, etc.
Mozi.m, [70, 71, 72, 73, 74, 75]	P2P ботнет, използващ DSHT	Септември, 2019	> 15 858	MIPS, RISC
Mukashi, [76, 77, 78]	Само-разпространяващ се червей	Февруари, 2020	-	ARM5-7, MIPS, MP5L, x86, etc.
Gafgyt, [1, 2, 3, 4, 114]	-	2014	> 1000000 (2016)	440FP, ARM4-7, i586, i686,

				M68K, MIPS, PowerPC, Sh4, SPARC, x86, etc.
Carna, [5, 6, 79, 80]	-	-	> 3000000	Unknown
Chuck Norris, [7, 8]	IRC бот с C&C	2009	> 33000	MIPSEL
LightAidra, Aidra, [9]	-	2012	>420000	ARM, MIPS, MIPSEL, PPC, x86/x86-64, SuperH
Linux.Darlloz [10, 11, 12, 13, 14]	Червей	Ноември, 2013	> 31000	x86, ARM, MIPS, PowerPC
Linux.Wifatch [15, 16, 17]	-	2014	>10000	ARM, MIPS, SH4 (mainly), PowerPC, x86 (minority)
Linux/Hydra [18, 19]	Активности, наподобяващи червей	2008	-	x86, MIPSEL
Linux/IRCTelnet [19, 20, 21, 22, 23, 24, 25, 26]	-	2016	-	ARM, PPC, SPH, i32, etc.
Psyb0t [36, 37, 38, 39]	Червей	2009	~100000	MIPS

Remaiten [19, 40, 41, 42]	-	2016	-	ARM, MIPS, PowerPC, SuperH
Spike-Dofloo [43, 44, 45, 46, 47, 48, 49, 50, 115]	-	2019	-	-
TheMoon [51, 52, 53, 54, 55, 117]	Червей (ver. 2 – троянски кон)	2014	~2000	MIPS
Tsunami [56, 57, 58, 59, 116]	Само-разпространяващ се червей	2010	> 600 000	ARM, MIPS, MPSL

Изводи

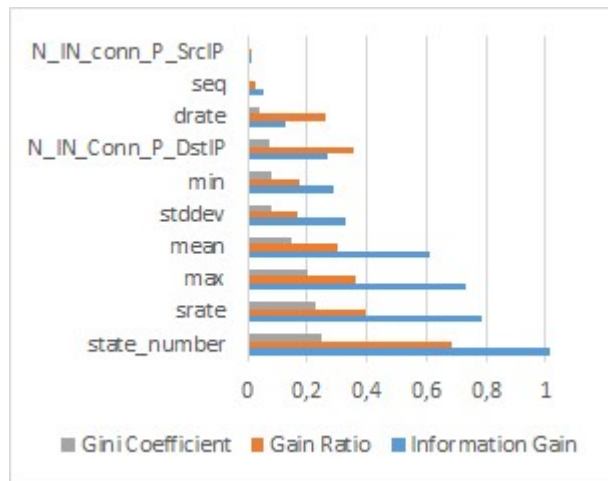
Въз основа на видовете атаки, извършени с използване на ботнет, в които значителна част от ботовете се състоят от IoT устройства, има няколко възможни стратегии, които да се приложат, за да се открият и допълнително ограничат злонамерените действия:

1. Мониторинг на приложния слой за аномално поведение на вградени програми.
2. Мониторинг на трафика към и извън IoT устройството.
3. Мониторинг на физическия слой.
4. Изграждане на бази със записи от обобщаващи признаци за детекция и разпознаване на осъществявани в реално време IoT базирани ботнет атаки.

ГЛАВА 2 - ОТКРИВАНЕ И КЛАСИФИКАЦИЯ НА БОТНЕТ АТАКИ ЧРЕЗ НЕВРОННИ МРЕЖИ С ОБРАТНО РАЗПРОСТРАНЕНИЕ НА ГРЕШКАТА

Работи се с тестовата база Bot-IoT [130]. Условните кодове на отделните атаки са 0 (без атака), 1 (DoS TCP атака), 2 (DoS UDP атака), 3 (DoS HTTP атака), 4 (DDoS TCP атака), 5 (DDoS UDP атака), 6 (DDoS HTTP атака), 7 (Keylogging), 8 (Data Exfiltration), 9 (OS Fingerprint) и 10 (Service Scan). Получените стойности за информационооно усливане IG , присъщата стойност IV и коефициента на Джини G за пълния набор от обучаващи отчети са показани на фиг. 2.2. Най-забележимото разграничение в разпределението на параметрите на характеристиките се отнася до seq и N_IN_conn_P_SrcIP, които са с 2 порядъка по-ниски от останалите 8 признака. Това е основният

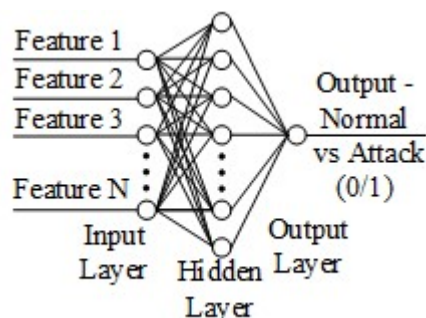
мотивиращ фактор в нашия експеримент да се опитаме да различим нормалния трафик от трафика при атака, като използваме не само всичките 10 признака, но и само 8 от тях.



Фиг. 2.2. Ранговане на признаците

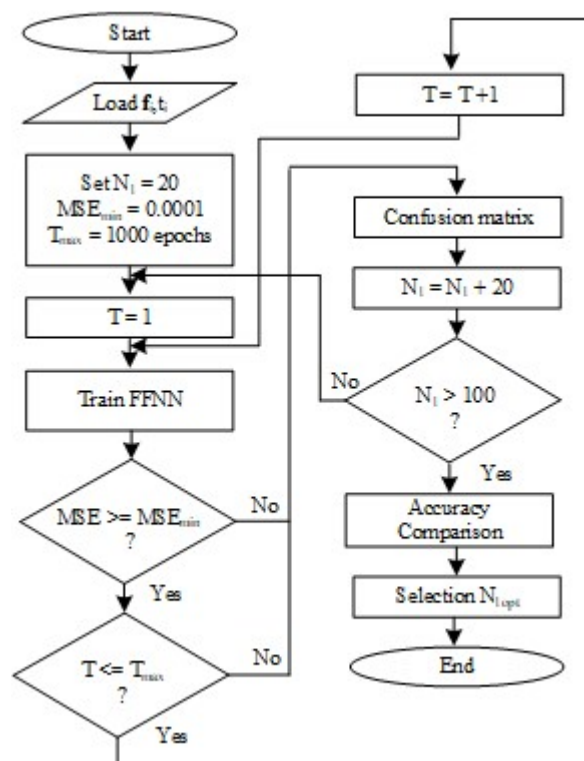
Структура на детектор с еднослойна невронна мрежа с обратно разпространение на грешката

Една от основните цели на представеното изследване е да предложи по-опростена структура на класификатор спрямо други, като например рекурентна невронна мрежа, описана в [130] за същата задача за разграничаване на атаките от нормалните мрежови активности чрез мониторинг на трафика. Това позволява, както времето за изпълнение по време на обучението и класификацията, така и консумацията на памет да бъдат пониски. Общата структура на класификатора, разработена тук, е представена на фиг. 2.5. Това е невронна мрежа с право разпространение на признаците и алгоритъм за обратно разпространение на грешката при обучение [149].



Фиг. 2.5. Предлагана невронна мрежа за детекция на мрежови атаки

Няма общ подход, който да посочи как да се избере броят на невроните в скрития слой N_I на мрежата, описан по-горе. Ето защо по време на експеримента се прилага следната обща схема (фиг. 2.6) за техния избор.



Фиг. 2.6. Оптимизация на невронната мрежа

Експериментални резултати - хардуерната платформа, използвана за тестване, се състои от процесор Intel Xeon E5-1620 с 4 ядра, работещ в режим на хипер-нишкова обработка на 3.50 GHz. Процесорът разполага с L1 кеш от 256 kB, L2 - 1 MB и L3 - 10 MB. Размерът на RAM паметта е 64 GB, а размерът на твърдия диск е 2 TB. Всички компоненти са под контрола на MS Windows 10 Pro 20H2. Orange v. 3.28 е симулационната среда, в която всички модели на невронни мрежи се обучават и тестват. Времената за изпълнение са дадени в табл. 2.2.

Табл. 2.2. Времена на изпълнение, постигати чрез SGD и Adam оптимизация при 80 неврона в скрития слой

Алгоритъм	Активационна функция	Training Time, sec	Testing Time, sec	
			Train set	Test set
SGD	Identity	138.79	2.59	0.95
	Logistic	159.50	7.86	1.97
	Tanh	243.18	7.82	2.19
	ReLU	228.59	5.93	1.59
Adam	Identity	145.88	3.01	0.95

	Logistic	179.65	8.05	9.17
	Tanh	212.47	11.53	3.36
	ReLU	189.19	4.61	1.89

Оптимизацията по алгоритъма на Adam се оказва много по-ефективна от SGD при този проблем с двоична класификация. Най-добрата точност е намерена за функцията на активиране от вида Tanh, която е допълнително тествана. Матриците на съвпадение при обработка на пълния набор от тестови отчети с оптималния класификатор, който се оказва, че има $N_{opt} = 60$ неврона за 8 признака и $N_{opt} = 80$ неврона за 10 признака, са дадени на фиг. 2.8. В табл. 2.9 е дадено сравнение с друга реализация на класификатор за същия тип данни.

Adam Tanh 10 features 60 neurons Binary output Test		Predicted		Σ	Adam Tanh 8 features 80 neurons Binary output Test		Predicted		Σ
		0	1				0	1	
Actual	0	98	9	107	Actual	0	80	27	107
	1	27	733571			1	28	733570	
Σ		125	733580	733705 _a	Σ		108	733597	733705 _б

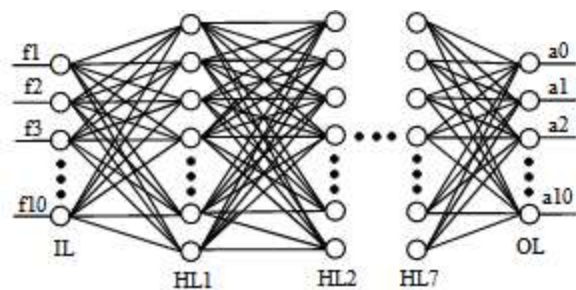
Фиг. 2.8. Матрици на съвпадението за оптимален детектор при: а – 10 признака и б – 8 признака

Табл. 2.9. Сравнение на класификационната ефективност на предложените детектори с RNN върху осреднени стойности на параметри за двата класа

Параметър	RNN - 10 признака, [1]	Предложен - 10 признака	Предложен - 8 признака
Precision	0.9999	0.9999	0.9998
Recall	0.9975	0.9999	0.9998
F1	0.7338	0.9999	0.9998
CA	0.9974	0.9999	0.9998

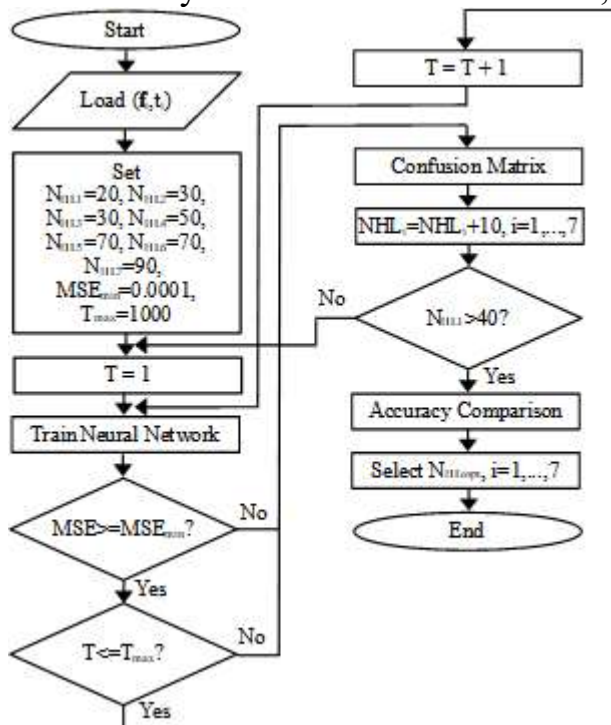
Структура на класификатор с многослойна невронна мрежа с обратно разпространение на грешката

Като се има предвид работата на Koroniotis и др. [130], където е предложена рекурентна невронна мрежа със 7 скрити слоя за класификация на множество мрежови атаки, базирани на IoT, предимно DDoS, в рамките на това изследване се предлага алтернативна структура на невронна мрежа. Това е невронна мрежа с право разпространение на данните със сходна степен на сложност по брой неврони, отново със 7 скрити слоя, но без обратни връзки в нея (фиг. 2.11).



Фиг. 2.11. Структура на многослойна невронна мрежа за класификация на IoT базирани мрежови атаки

Не съществува точно правило за избор на броя на невроните в скритите слоеве на невронна мрежа с обратно разпространение на грешката, в този случай $N_{HL2}, N_{HL3}, N_{HL4}, N_{HL5}, N_{HL6}, N_{HL7}, N_{HL8}$. В настоящото изследване е избран итеративен подход с поетапно увеличаване на тези числа, съгласно фиг. 2.12.



Фиг. 2.12. Алгоритъм за оптимизиране на структурата на многослойната невронна мрежа

Както е описано в анализа по-долу, оптималната конфигурация за 10 признака е избрана да бъде Adam 10-30-40-40-60-80-80-100-11, а за 8 признака – Adam 10-30-40-40-60-80-80-100-11. Матриците на съвпадение от класификацията над тестовия набор са показани съответно на фиг. 2.13 и фиг. 2.14. В табл. 2.17-2.19 е дадено сравнение по отношение на точността с друга реализация на класификатор за същия вид атаки.

		Predicted										
		0	1	2	3	4	5	6	7	9	10	Σ
Actual	0	98	0	2	0	0	0	1	0	2	4	107
	1	1	117279	0	1	5683	0	0	0	32	189	123185
	2	2	2	206571	0	1	50	0	0	0	0	206626
	3	1	6	0	272	5	0	12	0	0	5	301
	4	0	7001	0	0	186662	1	1	0	325	1162	195152
	5	0	0	206	0	2	189746	0	0	0	0	189954
	6	0	0	0	58	4	0	139	0	0	2	203
	7	5	0	0	0	0	0	0	8	0	1	14
	9	6	20	0	1	93	0	0	0	2084	1417	3621
	10	18	65	0	27	13	0	6	0	590	13823	14542
	Σ	131	124373	206779	359	192463	189797	159	8	3033	16603	733705

Фиг. 2.13. Матрица на съпадението за мрежа с 10 признака от класификация върху тестовата извадка

Вирну тестовата ковадка												
		Predicted										
		0	1	2	3	4	5	6	7	9	10	Σ
Actual	0	88	2	4	0	2	0	0	0	3	8	107
	1	1	88393	1	13	34743	0	29	0	0	5	123185
	2	0	3	206329	0	3	291	0	0	0	0	206626
	3	1	76	0	99	28	0	92	0	0	5	301
	4	2	3582	0	2	191527	0	4	0	0	35	195152
	5	3	3	1093	0	0	188851	0	0	4	0	189954
	6	0	23	0	2	52	0	124	0	0	2	203
	7	0	8	0	0	0	3	0	2	0	1	14
	9	7	78	0	5	90	0	2	0	168	3271	3621
	10	19	381	0	0	79	0	0	0	53	14010	14542
	Σ	121	92549	207427	121	226524	189145	251	2	228	17337	733705

Фиг. 2.14. Матрица на съпадението за мрежа с 8 признака от класификация върху тестовата извадка

Таблица 2.17. Сравнение на точността между RNN [130] и предложения класификатор за DDoS HTTP, TCP и UDP

Параметър	CA	Precision	Recall	F1
DDoS HTTP, RNN [1]	0.9932	0.9930	0.9970	0.0147

DDoS HTTP, предложен	0.9998	0.8742	0.6847	0.7679
DDoS TCP, RNN [1]	0.9999	0.9999	0.9999	0.0105
DDoS TCP, предложен	0.9805	0.9698	0.9565	0.9631
DDoS UDP, RNN [1]	0.9999	0.9999	0.9999	0.0063
DDoS UDP, предложен	0.9996	0.9997	0.9989	0.9993

Табл. 2.18. Сравнение между точността на RNN [130] и предложения класификатор за DoS HTTP, TCP и UDP

Параметър	<i>CA</i>	<i>Precision</i>	<i>Recall</i>	<i>F1</i>
DoS HTTP, RNN [1]	0.9806	0.9898	0.9845	0.0314
DoS HTTP, предложен	0.9998	0.7576	0.9036	0.8242
DoS TCP, RNN [1]	0.9999	0.9999	0.9999	0.0713
DoS TCP, предложен	0.9822	0.9429	0.9521	0.9475
DoS UDP, RNN [1]	0.9999	0.9999	0.9999	0.0126
DoS UDP, предложен	0.9996	0.9990	0.9997	0.9993

Табл. 2.19. Сравнение на точността на класификация при RNN [130] и предложената многослойна невронна мрежа за OS Fingerprinting, Service Scan, Data Exfiltration и Keylogging

Параметър	<i>CA</i>	<i>Precision</i>	<i>Recall</i>	<i>F1</i>
OS Finger, RNN [1]	0.9917	0.9984	0.9931	0.0587
OS Finger, предложен	0.9952	0.8325	0.9505	0.8876
Service Scan, RNN [1]	0.9957	0.9986	0.9971	0.2159
Service Scan, предложен	0.9952	0.8342	0.9515	0.8890
Data Exfilt., RNN [1]	0.9875	0.0000	0.0000	0.0000
Data Exfilt., предложен	0.9966	0.6871	0.5755	0.6264
Keylogging, RNN [1]	0.9873	0.9853	0.9178	0.0021
Keylogging, предложен	0.9999	1.0000	0.5714	0.7272

Изводи

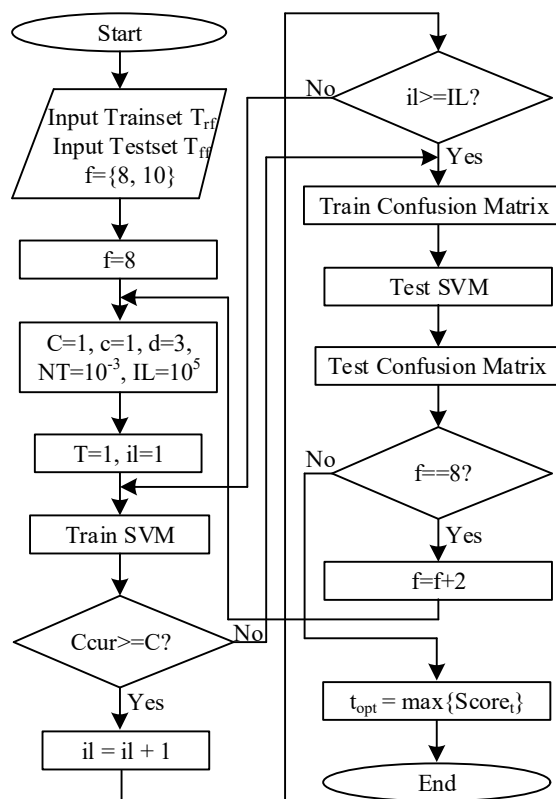
В настоящата глава е предложен нов модел на невронна мрежа с право предаване на данните и обратно разпространение на грешката с един скрит слой, способен да открива DoS и DDoS атаки, както и кражба на информация и проучвателен тип активности през мрежи за комутиране на IP пакети, насочени към конкретни машини-жертви. Оптималната структура на класификатора се състои от 10 входа, 60 скрити неврона, използващи функция за активиране хиперболичен тангенс и единичен изход за двоична дискриминация за атака срещу активност без атака. Най-точен резултат от обучението се постига чрез процедурата за оптимизация на Adam.

Аналогичният модел на невронна мрежа, включващ само 8 от 10 признака с 80 скрити неврона, работи също достатъчно точно и заедно с изпълнението с 10 признака постига по-добра точност от рекурентна невронна мрежа с 4 скрити слоя при една и съща база със записи за тестване. Предложеният класификатор се счита за обещаващ за практическо внедряване като компонент на онлайн система за наблюдение, главно срещу DDoS атаки в мрежи с комутация на пакети.

Също така в Глава 2 е предложена нова подобрена многослойна невронна мрежа с обратно разпространение на грешката за класификация на базирани на IoT DoS, DDoS и други злонамерени мрежови дейности, която може да бъде внедрена в практиката. Голямо разнообразие от тестове доказват, че оптимизационната процедура на Adam е по-ефективна от алгоритъма Scaled Gradient Descend (SGD). Активационната функция тангенс хиперболичен е подходяща за всички неврони от скритите слоеве. Времето за обучение и класификация предполагат приложимостта на невронния модел в реални приложения. Използването на 8, по-малко информативни признаци, от пълния набор от 10 признака, изглежда балансирано решение по отношение на точността на класификацията, което би могло да ускори процеса на обучение поне два пъти. Някои от атаките, като Keylogging и OS Fingerprinting, предвид значително по-ниския интензитет на генерирания трафик, са по-трудни за откриване с около 57%. Необходими са допълнителни модификации на класификатора, евентуално включването му в съставна реализация с други класификатори, за да се регистрират подобни дейности, еднакво добре с по-интензивните атаки.

ГЛАВА 3 - ОТКРИВАНЕ И КЛАСИФИКАЦИЯ НА БОТНЕТ АТАКИ ЧРЕЗ МАШИНИ С ПОДДЪРЖАЩИ ВЕКТОРИ

Общата процедура за оптимизация, която се предлага тук, е представена на фиг. 3.3. Съгласно препоръките, дадени в [164], началната стойност за параметъра Cost на SVM е $C = 1$, постоянният член $c = 1$, степента $d = 3$, числовият толеранс $NT = 10^{-3}$, границата на итерацията $IL = 10^5$. Един от основните ефекти от обучението е намирането на оптималната стойност за g . Като емпирично правило [164], началната стойност за него може да бъде зададена като $1/k$, където k е броят на компонентите на векторите на характеристиките, който е 0,1 в нашия случай. Въведена е следната категорийна променлива, обозначаваща типа на ядрото на SVM - $t = 1$ за линейно, $t = 2$ - за полином, $t = 3$ - за RBF и $t = 4$ - за сигмоидална. Текущата итерация по време на обучение се обозначава с il .



Фиг. 3.3. Процедура за оптимизация на SVM

Времената за обработка при детекция и класификация на атаките са дадени в табл. 3.2 и табл. 3.3.

Табл. 3.2. Времена за обработка с SVM за откриване на зловреден мрежови трафик спрямо нормален трафик

SVM тип	Брой признаци	Training time, s	Validation time, s	Testing time, s
Linear	8	8 776	68	15
	10	11 760	134	34
Polynomial	8	21 341	85	26
	10	12 677	143	36
RBF	8	76 539	228	57
	10	71 458	272	70
Sigmoid	8	137 980	218	54
	10	23 425	145	32

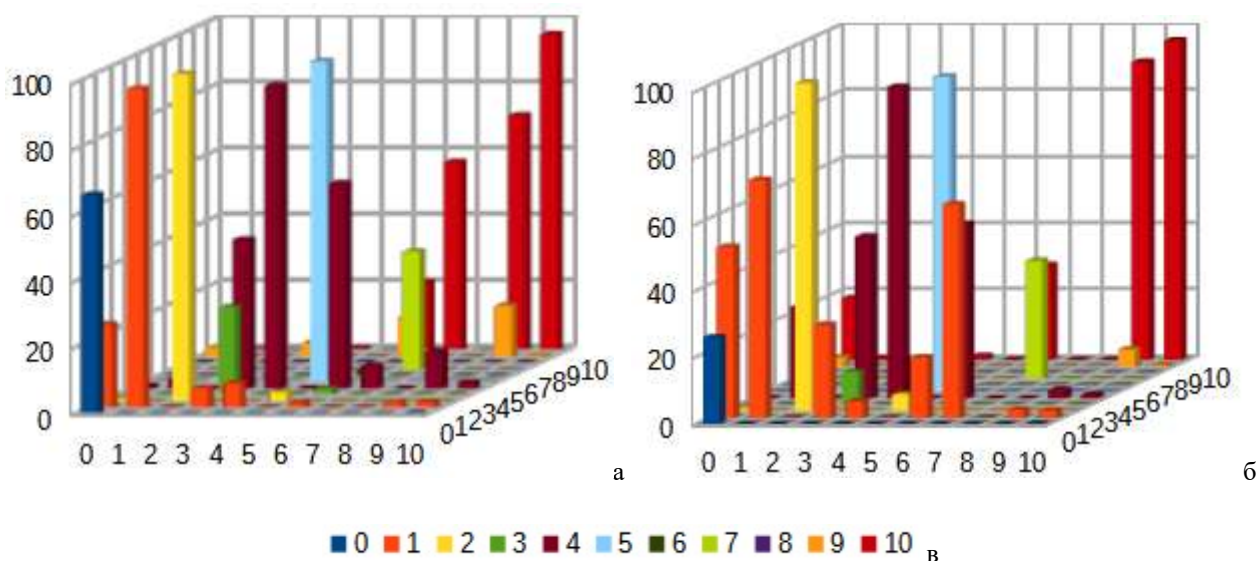
Табл. 3.3. Времена за обработка с SVM при класифициране на различни видове атаки

SVM тип	Брой итерации	Training time, s	Validation time, s	Testing time, s
RBF	8	1124612	55 543	32 558
	10	508 111	133 136	33 237

Класификационната точност за двата случая е дадена в табл. 3.4 и фиг. 3.4.

Табл. 3.4. Открити атаки като част от действителните атаки

SVM тип	Брой признаци	Неатаки, %		Атаки, %	
		Train	Test	Train	Test
Linear	8	2.7	3.7	100.0	100.0
	10	2.7	3.7	100.0	100.0
Polynomial	8	10.5	8.4	100.0	100.0
	10	33.8	29.9	100.0	100.0
RBF	8	11.9	11.2	100.0	100.0
	10	35.7	33.6	100.0	100.0
Sigmoid	8	0.8	0.9	100.0	100.0
	10	2.2	1.9	100.0	100.0



Фиг. 3.4. Матрици на съвпадението от класификация върху тестовия набор с:
а – 10 признака, б – 8 признака, в – цветова легенда на класовете

Сравнение с точността, постигана от друг класификатор от подобен вид, е представено в табл. 3.14 и табл. 3.15.

Табл. 3.14. Матрици на съвпадението на сравняваните SVM детектори

Предложен тук, в %			Предложен в [11], в %		
Атака	0	1	Атака	0	1
0	33.6	66.7	0	100.0	0
1	0.0	100.0	1	11.63	88.37

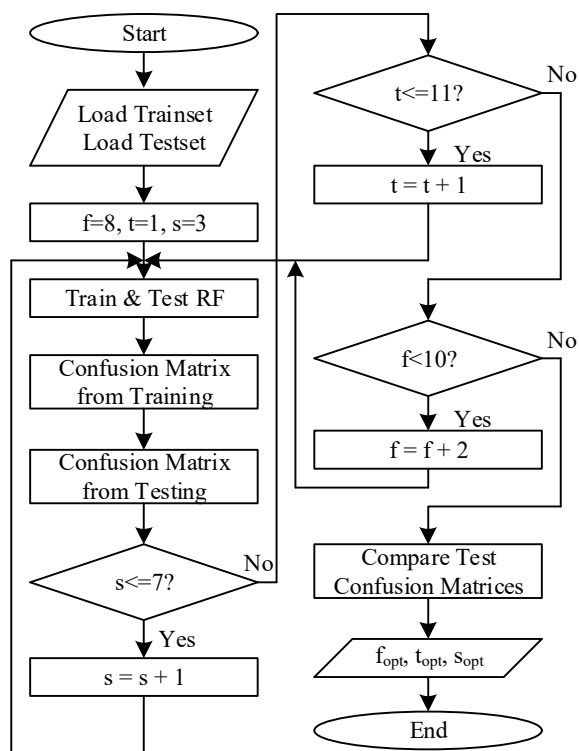
Табл. 3.15. Параметри за оценка на сравняваните SVM детектори

Параметър	Предложен тук	Предложен в [11]
Accuracy	0.9999	0.8837
Precision	0.9998	1
Recall	0.9999	0.8837
F1-measure	0.9998	0

Изводи - в тази част от изследванията са представени два вида SVM класификатори за DDoS атаки - двоични и многокласови, обхващащи 10 вида най-популярни злонамерени дейности, извършвани често с помощта на IoT ботнети. Най-точният детектор използва RBF ядро, което също се смята за най-правилното решение за многокласовата класификация. Най-бързият SVM класификатор е линейният, следван от полиномният, сигмоидален и RBF накрая, което важи и за повечето от случаите на обучение и тестване. Обучението с 8 признака се оказва по-бавно в повечето случаи, отколкото с 10 признака, но тестването може да бъде значително по-бързо с 8 признака за някои от ядрата, използвани в SVM. В класификацията на много класове, 10-те признака предлагат забележимо по-висока точност, отколкото 8-те признака. Този ефект е по-слабо изразен в двоичната класификация, но все още е валиден като обща тенденция. Предвид постигнатата точност, както бинарните, така и многокласовите SVM, представени в това изследване в техните оптимални конфигурации, се смятат за приложими в реални системи за мониторинг срещу DDoS атаки. Все пак е необходима допълнителна работа, особено за повишаване на точността на двоичния SVM при откриване на отчети без атака на фона на продължаваща DDoS атака с характерния ѝ висок интензитет. Могат да се приложат различни стратегии, които включват свързани класификатори, интелигентна селекция на отчети от набора за обучение и др.

ГЛАВА 4 - ОТКРИВАНЕ И КЛАСИФИКАЦИЯ НА БОТНЕТ АТАКИ ЧРЕЗ АНСАМБЛИ ОТ ДЪРВЕТА ЗА ВЗЕМАНЕ НА РЕШЕНИЯ

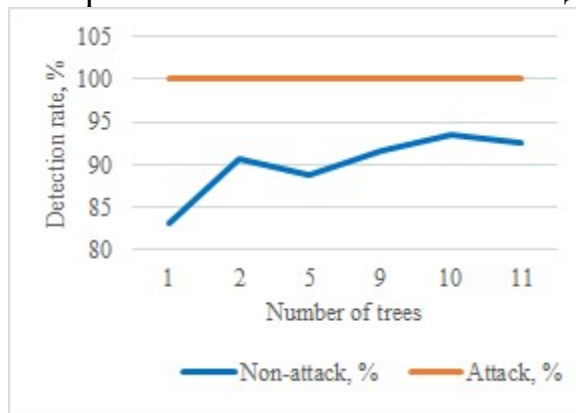
В представената в настоящото изследване реализация на RF алгоритъма има два регулируеми параметъра, класифициращи DDoS атаки от агрегирани мрежови записи (фиг. 4.3). Първият е броят на дърветата t , който обикновено е около 10 [189], а в текущото проучване той варира между 1 и 11. Другият параметър е минималният брой подмножества s , които не трябва да се разделят допълнително по време на обучение и тестване. Тази променлива, която настройваме е от 3 до 7, чийто диапазон, както ще видим по-долу, причинява известна промяна в степента на откриване на различни атаки. Цялото тестване е направено с 10 и 8 компоненти за отчетите, за което параметърът f е включен в диаграмата на предлагания алгоритъм за оптимизация (фиг. 4.3).



Фиг. 4.3. Оптимизационна процедура за RF

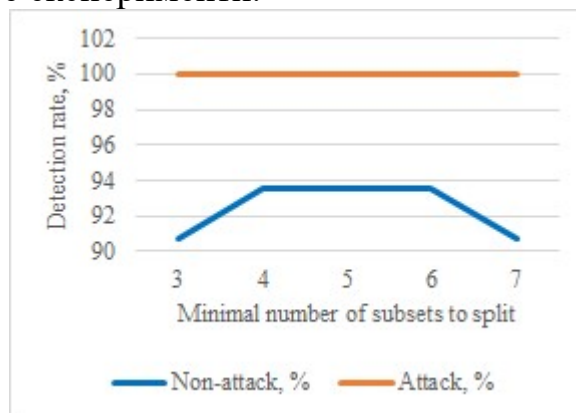
Изпълнението на процедурата за оптимизация, предложена в т. 4.4 (фиг. 4.3), доведе до резултатите, представени на фиг. 4.4 по отношение на броя на изпробваните дървета при извършване на процеса на детекция (двоична класификация). Частта от случаите, съответстваща на коректна детекция на атаките, остава постоянна и почти равна на 100% в почти всички случаи, когато броят на дърветата варира между 1 и 11. В същото време има максимум при 10

дървета, достигайки точност на детекция около 94% за отчетите без атака. Поради това $t_{opt} = 10$ е избраната стойност за всички следващи експерименти.



Фиг. 4.4. Честота на коректна детекция на DDoS атаки във функция на броя дървета

Зависимостта на честотата на коректна детекция като функция на минималния брой подмножества за разделяне (фиг. 4.5) предлага малко по-различна картина от тази на общия брой дървета като независим параметър. Отново, честотата на детекция е постоянна и почти равна на 100% в почти всички случаи, но успехът при откриването на отчети без атака преминава в ниво на насищане за интервала между 4 и 6 подмножества - близо до 94%. Преди и след този интервал има спад в стойността на този параметър, който намалява до почти 90%. Направена е селекция в центъра на този интервал на насищане, която съответства на $s_{opt} = 5$ и е предпочитана при по-нататъшно тестване в останалите експерименти.



Фиг. 4.5. Честота на коректна детекция на DDoS атаки във функция на минималния брой подмножества за разделяне

Матриците на съвпадение от откриване на атаки над тестовия набор, използващи 10 и 8 компоненти на отчет, са графично представени съответно на фиг. 4.6 а и б. Времената за детекция и класификация са дадени в табл. 4.11.

Actual	Predicted		Σ
	0	1	
0	100	7	107
1	1	733597	733598
Σ	101	733604	733705

а

Actual	Predicted		Σ
	0	1	
0	103	4	107
1	2	733596	733598
Σ	105	733600	733705

б

Фиг. 4.6. Матрици на съвпадението при детекция на DDoS атаки върху тестовия набор при: а – 10 признака, б – 8 признака

Табл. 4.11. Времена за обработка от предлаганите RF класификатори

Режим на RF	Брой признаци	Training time, s	Validation time, s	Testing time, s
Детекция	8	179.94	15.00	6.21
	10	335.88	31.24	7.87
Класификация	8	311.35	54.94	13.63
	10	399.77	55.14	12.69

Матриците на съвпадение при тестване с 8 и 10 признака са дадени на фиг. 4.7 и фиг. 4.8.

Actual	Predicted										Σ
	0	1	2	3	4	5	6	7	9	10	
0	98	0	2	0	0	0	0	0	1	6	107
1	1	123180	0	1	3	0	0	0	0	0	123185
2	0	4	206619	2	0	1	0	0	0	0	206626
3	1	0	0	300	0	0	0	0	0	0	301
4	0	4	1	0	195147	0	0	0	0	0	195152
5	0	0	1	0	0	189952	0	0	0	1	189954
6	0	0	0	0	1	0	201	0	1	0	203
7	0	0	0	0	0	0	0	13	0	1	14
9	0	1	0	0	0	0	0	0	3304	316	3621
10	1	0	0	2	0	0	0	0	232	14307	14542
Σ	101	123189	206623	305	195151	189953	201	13	3538	14631	733705

Фиг. 4.7. Матрица на съвпадението за RF от тестовия набор с DDoS атаки при 10 признака

Actual	Predicted										Σ
	0	1	2	3	4	5	6	7	9	10	
0	101	1	2	0	1	1	0	0	1	0	107
1	1	122811	0	1	372	0	0	0	0	0	123185
2	1	3	206618	1	0	2	0	0	1	0	206626
3	0	0	0	299	1	0	1	0	0	0	301
4	0	460	0	0	194691	0	0	0	1	0	195152
5	0	0	0	0	0	189951	0	0	1	2	189954
6	0	1	0	1	1	0	199	0	0	1	203
7	1	0	0	0	0	0	0	12	0	1	14
9	0	1	0	0	84	0	0	0	930	2606	3621
10	1	0	0	1	65	0	0	0	797	13678	14542
Σ	105	123277	206620	303	195215	189954	200	12	1731	16288	733705

Фиг. 4.8. Матрица на съпадението за RF от тестовия набор с DDoS атаки при 8 признака

Сравнение на работата на предложения в тази глава класификатор с друг подобен на него е представено в табл. 4.12-4.13.

Табл. 4.12. Сравнение между два RF модела в режим на детекция

Параметър	RF, [191]	Предложен RF - 10 признака	Предложен RF - 8 признака
<i>CA</i>	0.9532	0.9999	0.9999
<i>Precision</i>	0.9580	0.9999	0.9999
<i>Recall</i>	0.9532	0.9999	0.9999
<i>F1-score</i>	0.9481	0.9999	0.9999

Табл. 4.13. Сравнение между два RF модела в режим на класификация

Параметър	RF, [191]	Предложен RF - 10 признака	Предложен RF - 8 признака
<i>CA</i>	0.9766	0.9992	0.9939
<i>Precision</i>	0.9824	0.9992	0.9932
<i>Recall</i>	0.9766	0.9992	0.9939
<i>F1-score</i>	0.9657	0.9992	0.9933

Изводи - В настоящата глава на дисертационния труд са предложени два нови модела на „случайни гори“, базирани на 8 и 10 компонентни характеристики от агрегирани записи на мрежови връзки, способни да откриват и класифицират 10 вида DDoS атаки, освен нормалния трафик. Няма значителна разлика по отношение на ефективността на откриване между 8- и 10-компонентните реализации на детектори. Първият е предпочитан за прилагане поради по-малкото количество данни за обработка и по-малкото време за обучение и тестване. Точността на класификацията се повишава значително при атаки за сканиране на услуги и атаки за сваляне на характеристики на ОС, когато се използват 10 признака от многокласовия RF класификатор. Това би било предпочитаният класификатор, който да се използва за пълния спектър от разглеждани атаки. Поради по-малките времена за обработка при 8 признака и по-малките количества обработвани данни, може да има практически случаи, когато този класификатор също би бил за предпочитане да се използва. Необходима е допълнителна работа, за да се намери оптималният набор от признаци, представляващи атаките за сваляне на характеристики на операционната система и извличане на използваните клавишни комбинации, където по-малко от 10 компонента биха довели до сравним, достатъчно висок процент на класификация за тях, както е за останалите 8 типа изследвани атаки.

ГЛАВА 5 - ОТКРИВАНЕ И КЛАСИФИКАЦИЯ НА БОТНЕТ АТАКИ ЧРЕЗ КОМБИНИРАНИ КЛАСИФИКАТОРИ

Описание на свързаните класификатори

Комбиниран двоичен класификатор

Точната форма на функцията на вероятността на комбинирания двоичен класификатор (детектор), предложена в този раздел, може да бъде изведена, както е показано по-долу, следвайки общия подход, даден в [203]. Изходът на единичните детектори NN, RF и SVM е съответно $o_1 \in \{0, 1\}$, $o_2 \in \{0, 1\}$ и $o_3 \in \{0, 1\}$. Стойността 0 означава отчет от нормален трафик, а 1 - атака. Изходът от модула за логистична регресия е $O \in \{0, 1\}$. Вероятността да се получи атака, предвид решението на комбинирания класификатор, е $p = P(O = 1)$. Според структурата на детектора от фиг. 5.1 лог-коэффициентите са:

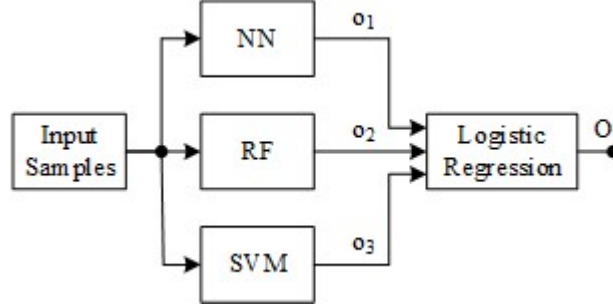
$$l = \ln \frac{p}{1-p} = k_0 + k_1 o_1 + k_2 o_2 + k_3 o_3, \quad (5.1)$$

където k_0 , k_1 , k_2 и k_3 са параметрите на модела за откриване на атаки. Повдигането на лявата и дясната страна на (5.1) на степен e води до:

$$\frac{p}{1-p} = e^{k_0 + k_1 o_1 + k_2 o_2 + k_3 o_3}, \quad (5.2)$$

което на свой ред дава:

$$p = \frac{e^{k_0+k_1o_1+k_2o_2+k_3o_3}}{e^{k_0+k_1o_1+k_2o_2+k_3o_3}+1} = \text{sigmoid}(e^{k_0+k_1o_1+k_2o_2+k_3o_3}). \quad (5.3)$$



Фиг. 5.1. Структура на предлагания класификатор

В обобщаваща форма при заложен линеен модел, следната параметризираща функция може да бъде дефинирана:

$$g_{\vec{k}}(\vec{o}) = \frac{1}{1+e^{-\vec{k}^T\vec{o}}} = P(O = 1|\vec{o}; \vec{k}), \quad (5.4)$$

където $\vec{k} = [k_0, k_1, k_2, k_3]$ и $\vec{o} = [1, o_1, o_2, o_3]$. Тогава, $P(O = 0|\vec{o}; \vec{k}) = 1 - g_{\vec{k}}(\vec{o})$ и вероятността за дадена реализация $o \in \{O\}$ се представя чрез:

$$P(o|\vec{o}; \vec{k}) = g_{\vec{k}}(\vec{o})^o (1 - g_{\vec{k}}(\vec{o}))^{(1-o)}. \quad (5.5)$$

Функцията на подобие, тогава, може да се изрази като:

$$\mathcal{L}(\vec{k}|o; \vec{o}) = P(O|\vec{o}; \vec{k}) = \prod_i P(o_i|\vec{o}_i; \vec{k}) = \prod_i g_{\vec{k}}(\vec{o}_i)^{o_i} (1 - g_{\vec{k}}(\vec{o}_i))^{(1-o_i)}. \quad (5.6)$$

Логаритъмът от $\mathcal{L}$ трябва да се максимизира чрез използване на алгоритъма с намаляване на градиента [204].

5.2.2. Комбиниран класификатор на множество атаки

Класифицирането на входните отчети в 10 типа атаки и такива за типичен трафик, т.е. $o_1, o_2, o_3, O \in \{0, 1, 2, \dots, 10\}$ може да бъде постигнато от единичните класификатори от фиг. 5.1 след комбинирането им с използването на логистична регресия от мултиномиален тип.

Следвайки общия подход от [203] и даден примерен вектор на резултатите от единичните класификатори $\vec{o}_i = [o_{1i}, o_{2i}, o_{3i}]$, получен като резултат от i -тото наблюдение през входа, функция за предсказване от линеен тип би могла да се дефинира, следвайки:

$$S(\vec{o}_i, c) = \vec{w}_c \vec{o}_i, \quad (5.7)$$

където c е текущ клас, $c \in \{C\} \equiv \{0, 1, 2, \dots, 10\}$ – един от възможните изходни резултати от класификацията, $\vec{w}_c$ – вектор на теглата, асоциирани с c -тия клас, а S е резултатът за нивото на принадлежност на $\vec{o}_i$ към c .

Следният израз, тогава, е валиден:

$$\left| \begin{array}{l} \ln \frac{P(O_i=1)}{P(O_i=0)} = \vec{w}_1 \vec{o}_i \\ \ln \frac{P(O_i=2)}{P(O_i=0)} = \vec{w}_2 \vec{o}_i \\ \vdots \\ \ln \frac{P(O_i=10)}{P(O_i=0)} = \vec{w}_{10} \vec{o}_i \end{array} \right. \quad (5.8)$$

Това е свързано с факта, че пълният мултиномиален модел с възможни изходни стойности на C може да бъде изграден върху двоични логистични регресии $C-1$, които са независими помежду си. Един от получените изходи е избран за база, а останалите $C-1$ се обработват, като се вземат като отправна точка. Ако за основа се избере $c = 0$, тогава се получава (5.8).

Повдигането на лявата и дясната страна на (5.8) на степен e в крайна сметка ще даде:

$$\left| \begin{array}{l} P(O_i = 1) = P(O_i = 0) e^{\vec{w}_1 \vec{o}_i} \\ P(O_i = 2) = P(O_i = 0) e^{\vec{w}_2 \vec{o}_i} \\ \vdots \\ P(O_i = 10) = P(O_i = 0) e^{\vec{w}_{10} \vec{o}_i} \end{array} \right. \quad (5.9)$$

и като се има предвид, че всички възможни изходни резултати образуват пълен набор, тогава:

$$P(O_i = 0) = 1 - \sum_{j=1}^{10} P(O_i = j) = 1 - \sum_{j=1}^{10} P(O_i = 0) e^{\vec{w}_j \vec{o}_i}. \quad (5.10)$$

От (5.10) е лесно да бъде намерено, че:

$$P(O_i = 0) = \frac{1}{1 + \sum_{j=1}^{10} P(O_i=0) e^{\vec{w}_j \vec{o}_i}}, \quad (5.11)$$

заедно със:

$$\begin{cases}
P(O_i = 1) = \frac{e^{\vec{w}_1 \vec{o}_i}}{1 + \sum_{j=1}^{10} P(O_i=0) e^{\vec{w}_j \vec{o}_i}} \\
P(O_i = 2) = \frac{e^{\vec{w}_2 \vec{o}_i}}{1 + \sum_{j=1}^{10} P(O_i=0) e^{\vec{w}_j \vec{o}_i}} \\
\vdots \\
P(O_i = 10) = \frac{e^{\vec{w}_{10} \vec{o}_i}}{1 + \sum_{j=1}^{10} P(O_i=0) e^{\vec{w}_j \vec{o}_i}}
\end{cases} \quad (5.12)$$

Компонентите на векторите $\vec{w}_j$, $j = 1, 2, \dots, 10$ могат да бъдат намерени чрез използване на метода максимум апостериори (maximum a posteriori) [205]. Ранговането на вероятностите от (11)-(12) за всеки входен набор $\vec{o}_i$ води до предсказания тип атака или нейната липса.

Експериментални резултати

Времената за обработка от свързаните класификатори са дадени в табл. 5.1 и 5.2. На фиг. 5.3 б и г са дадени матриците на съвпадение при тестване с 8 и 10 признака.

Табл. 5.1. Времена за обработка от детекторите

Детектор	Признаци	Training Time, sec	Validation Time, sec	Test time, sec
NN+RF	8	1541.12	16.82	4.66
	10	1914.30	23.22	5.86
NN+RF+SVM	8	412175.32	316.97	79.16
	10	197624.27	337.28	84.24

Табл. 5.2. Времена за обработка от класификаторите

Класификатор	Признаци	Training Time, sec	Validation Time, sec	Test time, sec
NN+RF	8	40809.83	106.83	31.78
	10	80890.32	84.87	29.70
NN+RF+SVM	8	123599.38	4935.23	1364.18
	10	195493.99	5020.27	1274.48

		Predicted											Σ
		0	1	2	3	4	5	6	7	8	9	10	
Actual	0	103	0	1	0	1	0	0	0	0	2	0	107
	1	1	122801	0	0	383	0	0	0	0	0	0	123185
	2	0	4	206620	1	0	0	0	0	0	1	0	206626
	3	0	1	0	296	1	0	1	0	0	1	1	301
	4	0	454	0	0	194697	0	0	0	0	1	0	195152
	5	0	0	2	0	0	189949	0	0	0	0	3	189954
	6	0	1	0	1	1	0	198	0	0	1	1	203
	7	1	0	0	0	0	1	0	8	0	0	4	14
	8	0	0	0	0	0	0	0	0	0	0	1	1
	9	0	0	0	0	84	0	0	0	0	510	3027	3621
	10	3	0	0	1	64	0	0	1	0	261	14212	14542
Σ		108	123261	206623	299	195231	189950	199	9	0	777	17249	733706

6

		Predicted											Σ
		0	1	2	3	4	5	6	7	8	9	10	
Actual	0	101	0	1	0	0	0	1	0	0	0	4	107
	1	0	123180	0	1	4	0	0	0	0	0	0	123185
	2	1	2	206622	1	0	0	0	0	0	0	0	206626
	3	0	0	0	298	1	0	2	0	0	0	0	301
	4	0	6	1	0	195145	0	0	0	0	0	0	195152
	5	0	0	2	0	0	189952	0	0	0	0	0	189954
	6	0	0	0	0	1	0	201	0	0	0	1	203
	7	0	0	0	0	0	1	0	12	0	0	1	14
	8	0	0	0	0	0	0	0	0	1	0	0	1
	9	0	0	0	0	0	0	0	0	0	3252	369	3621
	10	1	1	0	1	0	0	0	0	0	187	14352	14542
Σ		103	123189	206626	301	195151	189953	204	12	1	3439	14727	733706

г

Фиг. 5.3. Матрици на съвпадението за NN+RF класификаторите от: б – тестване при 8 признака, г – тестване при 10 признака

Фиг. 5.4 б и г изобразяват матриците на съвпаденията от класификация с

		Predicted											
		0	1	2	3	4	5	6	7	8	9	10	Σ
Actual	0	102	1	2	0	1	0	0	0	0	1	0	107
	1	1	122854	0	1	329	0	0	0	0	0	0	123185
	2	0	3	206621	1	0	0	0	0	0	1	0	206626
	3	0	0	0	297	2	0	2	0	0	0	0	301
	4	0	472	0	0	194679	0	0	0	0	1	0	195152
	5	0	0	2	0	0	189949	0	0	0	1	2	189954
	6	0	1	0	1	3	0	196	0	0	1	1	203
	7	0	0	0	0	0	1	0	9	0	4	0	14
	8	0	0	0	0	1	0	0	0	0	0	0	1
	9	0	0	0	0	84	0	0	0	0	483	3054	3621
	10	1	0	0	1	64	0	0	1	0	239	14236	14542
Σ		104	123331	206625	301	195163	189950	198	10	0	731	17293	733706

б

		Predicted											
		0	1	2	3	4	5	6	7	8	9	10	Σ
Actual	0	102	0	1	0	0	0	0	0	0	2	2	107
	1	1	123180	0	0	4	0	0	0	0	0	0	123185
	2	2	4	206619	0	0	1	0	0	0	0	0	206626
	3	1	0	0	298	1	0	1	0	0	0	0	301
	4	0	5	1	0	195146	0	0	0	0	0	0	195152
	5	0	0	2	0	0	189952	0	0	0	0	0	189954
	6	0	0	0	0	0	0	202	0	0	0	1	203
	7	0	0	0	0	0	1	0	12	0	0	1	14
	8	0	0	0	0	0	0	0	0	1	0	0	1
	9	0	1	0	0	0	0	0	0	0	3233	387	3621
	10	1	0	0	1	0	0	0	0	0	208	14332	14542
Σ		107	123190	206623	299	195151	189954	203	12	1	3443	14723	733706

г

Фиг. 5.4. Матрици на съвпадението за NN+RF+SVM при класификация от: б – тестване при 8 признака, г – тестване при 10 признака

Сравнение на точността с други класификатори на същия вид атаки е представено в табл. 5.16.

Табл. 5.16. Сравнение на точността при детекция на атаки

Детектор	SVM, [16], 10 признака	RNN, [16], 10 признака	LSTM, [16] 10 признака	Proposed NN+RF, 8 признака
CA	0.8837	0.9974	0.9974	0.9999
Precision	1.0000	0.9999	0.9999	0.9999
Recall	0.8837	0.9975	0.9975	0.9999

Изводи

В тази Глава са предложени 4 нови модела на двоични класификатори на мрежови атаки, базирани на използването на IoT устройства. Първият е детектор с 2 модула, различаващ отчети от нормален трафик от отчети с 10 типа атаки (DoS и DDoS разновидности на TCP, UDP и HTTP flood, Keylogging, Data Exfiltration, OS Fingerprinting и Service Scanning). Състои се от невронна мрежа с право разпространение с 1 скрит слой и RF, свързани чрез логистична регресия. Вторият е детектор с 3 модула - с допълнително въведена SVM. Една двойка детектори оперират с 8 признака, а другата двойка - с 10 признака - с 2 нови различни компонента от първия. Честотата на грешките на 2-модулния детектор при 8 признака варира в интервала $1.36 \cdot 10^{-3}\%$ и в $1.50 \cdot 10^{-3}\%$ при 10 признака. Детекторът с 3 модула има ниво на грешки в интервала $1.50 \cdot 10^{-3}\%$ при 8 признака и $1.91 \cdot 10^{-3}\%$ - при 10 признака. Грешките са концентрирани най-вече в отчетите без атака. Единичната невронна мрежа, както показва предишно проучване (Глава 2), има ниво на грешка от $4.91 \cdot 10^{-3}\%$ в оптималната си конфигурация като детектор, работещ с 10 признака. От други 2 предходни проучвания се намира, че SVM въвежда $9.81 \cdot 10^{-3}\%$ грешки, докато RF има $0.82 \cdot 10^{-3}\%$ грешки, като работи с 8 признака като детектор. Грешките, предизвикани от единичните детектори, са по-равномерно разпръснати между отчети с атака и от нормален трафик. Времето за обучение на детектора с 3 модула, в сравнение с оптималния 2-модулен вариант, при използване на 8 признака отнема 267,5 пъти повече време. Съотношението на времето за тестване между същите 2 конфигурации е почти 17 пъти. Детекторът с 2 модула с 8 признака се оказва оптималният детектор сред всички тествани комбинирани двоични класификатори в това изследване.

Научни, научно-приложни и приложни приноси

Научни приноси

1. Разработен е теоретичен модел на комбиниран двоичен класификатор на IoT-базирани мрежови атаки, включващ невронна мрежа с право разпространение, машина с поддържащи вектори и ансамбъл от дървета за

вземане на решение, които водят до по-точна детекция на атаките от представени в практиката решения ((5.1)-(5.6), фиг. 5.1).

2. Разработен е теоретичен модел на комбиниран класификатор на 10 вида IoT-базирани мрежови атаки и агрегирани записи от нормален трафик, включващ невронна мрежа с право разпространение на грешката, машина с поддържащи вектори и ансамбъл от дървета за вземане на решения, които водят до по-точна класификация от представени в практиката решения ((5.7)-(5.15), фиг. 5.1).

Научно-приложни приноси

1. Разработен е и е оптимизиран модел на еднослойна невронна мрежа с право разпространение на данните и обратно разпространение на грешката, която изпълнява роля на детектор и работи с 8 или 10 признака, представящи обобщение на мрежови трафик, съдържащ 10 вида IoT базирани ботнет атаки и записи с нормален трафик, която в първия случай има 80, а във втория – 60 неврона с тангенс хиперболичен като активационна функция и се обучава по оптимизационния алгоритъм на Adam като постига по-висока точност на детекция от по-висока по сложност рекурентна невронна мрежа, приложена вече в практиката ((2.8)-(2.11), фиг. 2.5-2.8, табл. 2.2-2.9).

2. Разработен е и е оптимизиран модел на многослойна невронна мрежа с право разпространение на данните и обратно разпространение на грешката, която изпълнява роля на класификатор и работи с 8 или 10 признака, представящи обобщение на мрежови трафик, съдържащ 10 вида IoT базирани ботнет атаки и записи с нормален трафик, която има 30-40-40-60-80-80-100 неврона в скритите си слоеве с тангенс хиперболичен като активационна функция и се обучава по оптимизационния алгоритъм на Adam като постига по-висока точност на класификация от еквивалентна по сложност рекурентна невронна мрежа, приложена вече в практиката ((2.17)-(2.21), фиг. 2.11-2.14, табл. 2.10-2.19).

3. Разработен е оптимизиран модел на машини с поддържащи вектори за детекция и класификация на 10 вида IoT ботнет базирани мрежови атаки, който се оказва по-точен от аналогични вече разработени модели в практиката (фиг. 3.3, табл. 3.2 – 3.13).

4. Разработен е оптимизиран модел на ансамбъл от дървета за взимане на решение за детекция и класификация на 10 вида IoT ботнет базирани мрежови атаки, който се оказва по-точен от аналогични вече разработени модели в практиката (фиг. 4.3-4.5, табл. 4.5 – 4.11)

Приложни приноси

1. Представен е модел на таксономия на IoT устройства, намиращи масова употреба в практиката (фиг. 1.1). Направена е оценка на разпространението на IoT базирани ботнет и свързаните с тях C&C сървъри, както и на финансовите загуби в различните индустрии, причинени от провеждането на разпределени мрежови атаки върху предлаганите от тях услуги (фиг. 1.2-1.6). Също така е направен анализ на основните видове архитектури и принцип на действие на съществуващите ботнет, чрез които се извършват разпределени мрежови атаки върху устройства и предлаганите от тях услуги (фиг. 1.7-1.11). Направен е и анализ на най-масово разпространените ботнет от последното десетилетие като е изяснен механизма на тяхното възникване, разширяване и провеждане на разпределени мрежови атаки срещу предлагани персонални или масов тип услуги, както и на последиците от тези атаки (табл. 1.1). Направени са изводи и са предложени конкретни стратегии и свързаните с тях програмни модели и средства от областта на машинното обучение с цел откриване на този вид злонамерени активности чрез анализ на мрежови трафик.

2. Изготвена е статистика на база редица представящи параметри на признаците, описващи 10 вида IoT базирани ботнет атаки и нормален трафик, които са използвани за обучение и тестване на еднослойни и многослойни невронни мрежи за тяхното откриване и класификация ((2.1)-(2.7), (2.12)-(2.16), фиг. 2.1-2.4 и фиг. 2.9-2.10, табл. 2.1). На база на тези признаци е извършено обучение, тестване и анализ на резултатите от работата на еднослойни и многослойни невронни мрежи с обратно разпространение на грешката за откриване на IoT базирани ботнет атаки (т. 2.2.4 и т. 2.3.3).

3. Изготвено е статистическо разпределение на вероятността за възникване на определен вид IoT базирана ботнет атака в зависимост от стойността на различни признаци, агрегирани като обобщаващи статистически характеристики от мрежови трафик (фиг. 3.1 и табл. 3.1). Извършено е обучение, тестване и анализ на резултатите от работата на машини с поддържащи вектори при 4 различни вида ядра за откриване на IoT базирани ботнет атаки (табл. 3.2 – 3.15, фиг. 3.4).

4. Намерена е пространствена трансформация за представяне на признаците, агрегирани от мрежови трафик, съдържащ връзки от IoT базирани ботнет атаки, която представя оптимално информативното значение на всеки признак (фиг. 4.1-4.2, табл. 4.3). Извършено е обучение, тестване и анализ на резултатите от работата на ансамбли от дървета за вземане на решение за откриване и класификация на IoT базирани ботнет атаки (табл. 4.5-4.13, фиг. 4.6-4.8).

5. Разработени са програмни реализации на комбинирани детектори и класификатори на IoT-базирани мрежови атаки, включващи невронна мрежа с право разпространение на грешката, машина с поддържащи вектори и ансамбъл от дървета за вземане на решения, които са обучени и тествани с общодостъпния тестови набор Bot-IoT като е доказана по-високата класификационна точност спрямо други решения от практиката и е оценено бързодействието им (табл. 5.1-5.17, фиг. 5.2-5.4).

Публикации по дисертационния труд

1. Ivanova, V., T. Tashev, I. Draganov, Survey on IoT based Network Attacks, In Proc. Of the XVIII-th International Conference "Challenges in Higher Education and Research in the 21st Century" (Cher'21), June 1-4, Sofia, Bulgaria, Virtual Conference, pp. 43-51, 2021.

2. Ivanova, V., T. Tashev, I. Draganov, Survey on IoT Botnets, In Proc. Of the XVIII-th International Conference "Challenges in Higher Education and Research in the 21st Century" (Cher'21), June 1-4, Sofia, Bulgaria, Virtual Conference, pp. 53-61, 2021.

3. Ivanova, V., T. Tashev, I. Draganov, Detection of IoT based DDoS Attacks by Network Traffic Analysis using Feedforward Neural Networks, International Journal of Circuits, Systems and Signal Processing, vol. 16, pp. 653-662, 2022.

4. Ivanova, V., Multiple IoT based Network Attacks Discrimination by Multilayer Feedforward Neural Networks, International Journal of Circuits, Systems and Signal Processing, vol. 16, pp. 675-685, 2022.

5. Ivanova, V., T. Tashev, I. Draganov, DDoS Attacks Classification using SVM, WSEAS Transactions on Information Science and Applications, vol. 19, pp. 1-11, 2022.

6. Ivanova, V., T. Tashev, I. Draganov, Random Forest Detector and Classifier of Multiple IoT-based DDoS Attacks, WSEAS Transactions on Information Science and Applications, 2022 (in press).

7. Ivanova, V., T. Tashev, I. Draganov, IoT-based Network Attacks Discovery with Combined Classifiers, International Journal of Circuits, Systems and Signal Processing, vol. 16, pp. 754-763, 2022.

SUMMARY



TECHNICAL UNIVERSITY OF SOFIA

FACULTY OF COMPUTER SYSTEMS AND TECHNOLOGIES

INFORMATION TECHNOLOGY IN INDUSTRY DEPARTMENT

Mag. Eng. Vanya Dimitrova Ivanova

Abstract of Ph.D. Thesis

IoT based Botnet Attacks Discovery by Network Traffic Analysis

In this PhD thesis, single and combined classifiers of multiple network attack types, including DoS TCP, DoS UDP, DoS HTTP, DDoS TCP, DDoS UDP, DDoS HTTP flood, Keylogging, Data Exfiltration, OS Fingerprint and Service Scan, initiated by IoT botnets, are proposed. All of them are implemented once as detectors of ongoing attacks, regardless of their type, discriminating them from normal traffic, and in another implementation - as multiclass classifiers, capable of recognizing each type of attack separately. The single classifiers are represented by feedforward neural networks, support vector machines and random forest. All of them seem applicable for discriminating the various kinds of attacks with the random forest being the most accurate and fast. In order to increase classification accuracy all three classifier types are combined in composite classifier and compared with a two-component classifier of only the feedforward neural network and random forest. The latter prove to be faster while demonstrating comparable accuracy with the case where the support vector machine also takes part in the classification process. All tested classifiers are considered applicable in the practice considering the ever-growing diversity of network attack types and the associated botnets with them. They are trained with 10- and 8-component aggregated from the network traffic samples where some of the classifiers cope at the same level of accuracy with just 8 of the features. It requires less memory for storage for the processed records and in some cases leads to faster execution times compared to the 10-element records. Further efforts are needed towards achieving higher accuracy of detection of the keylogging, data exfiltration, OS fingerprint and service scan since these activities have lower intensity than most of the DoS and DDoS attacks that generate much more samples for training and classifying and are more easily distinguished one from the other.